
Generative Decoding of Visual Stimuli

Eleni Miliotou^{* 1} Panagiotis Kyriakis^{* 2} Jason D Hinman¹ Andrei Irimia³ Paul Bogdan²

Abstract

Reconstructing real-world images from fMRI recordings is a challenging task of great importance in neuroscience. The current architectures are bottlenecked because they fail to effectively capture the hierarchical processing of visual stimuli that takes place in the human brain. Motivated by that fact, we introduce a novel neural network architecture for the problem of neural decoding. Our architecture uses Hierarchical Variational Autoencoders (HVAEs) to learn meaningful representations of real-world images and leverages their latent space hierarchy to learn voxel-to-image mappings. By mapping the early stages of the visual pathway to the first set of latent variables and the higher visual cortex areas to the deeper layers in the latent hierarchy, we are able to construct a latent variable neural decoding model that replicates the hierarchical visual information processing. Our model achieves better reconstructions compared to the state of the art and our ablation study indicates that the hierarchical structure of the latent space is responsible for that performance.

1. Introduction

Decoding visual imagery from brain recordings is a key problem in neuroscience. This problem aims to reconstruct the visual stimuli from fMRI recordings taken while the subject is viewing the stimuli. Even though some of the excitement is fuelled by science fiction and the difficulty of the problem (Shen et al., 2019b), the scientific consensus is that neural decoding has real-world, important implications. It is important for understanding how neural

activity relates to external stimuli (Glaser et al., 2020), for engineering application such as brain-computer interfaces (Brandman et al., 2017), for decoding imagery during sleep (Horikawa et al., 2013), and even for reconstructing video frames from brain activity (Le et al., 2022). Given its importance, neuroscience and machine learning researchers have jointly led the development of sophisticated deep learning architectures that allows us to design pipelines that map voxel-based recordings to the corresponding visual stimuli. Based on the target learning task, visual decoding can be categorized into *stimuli classification*, *stimuli identification*, and *stimuli reconstruction*. The former two tasks aim to predict the object category of the presented stimulus or identify the stimulus from an ensemble of possible stimuli. The reconstruction task, which is the most challenging one and the main focus of this paper, aims to construct a replica of the presented stimulus image from the fMRI recordings.

Related Work. The proposed methods for the problem of neural decoding can be broadly classified in three categories: *non-deep learning methods*, *non-generative deep learning methods* and *generative deep learning methods*. The *non-deep learning* class consists of methods that are based on primitive linear models and aim in reconstructing low-level image features (Kay et al., 2008). Such approaches first extract handcrafted features from real-world images, such as multi-scale image bases (Miyawaki et al., 2008) or Gabor filters (Yoshida & Ohki, 2020), and then learn a linear mapping from the fMRI voxel space to the extracted features. Due to their simplicity, linear models are not able to reconstruct complex real-world images and thus their applicability is restricted to simple images containing only low-level features.

Methods that use convolutional neural networks as well as encoder-decoder architectures belong to the *non-generative deep learning* class. Horikawa et al. (Horikawa & Kamitani, 2017) demonstrated a homology between human and machine vision by designing an architecture with which the features extracted from convolutional neural networks can be predicted from fMRI signals. Based upon those findings, Shen et al. (Shen et al., 2019b) used a pretrained VGG-19 model to extract hierarchical features from stimuli images and learned a mapping from the fMRI voxels in the low/high area to the corresponding low/high VGG-19 features. Belyi et al. (Belyi et al., 2019) designed a

^{*}Equal contribution ¹Department of Neurology, University of California Los Angeles, Los Angeles, US ²Department of Electrical and Computer Engineering, University of Southern California, Los Angeles, US ³Department of Gerontology, University of Southern California, Los Angeles. Correspondence to: Panagiotis Kyriakis <pkryriaki@usc.edu>.

CNN-based Encoder-Decoder architecture, where the encoder learns a mapping from the stimulus images to the fMRI voxels and the decoder learns the reverse mapping. By stacking the components back-to-back, the authors train their network using self-supervision, thereby addressing the inherent scarcity of fMRI-image pairs. Following up on that work, Gaziv et al. (Gaziv et al., 2020) improved the reconstruction quality by training on a perceptual similarity loss function, which is calculated by first extracting multi-layer features from both the original and reconstructed images and comparing the extracted features layer-wise. Such a perceptual loss is known to be highly effective in assessing the image similarity and accounts for many nuances in the human vision (Zhang et al., 2018).

In the *generative deep learning* class, we have model architectures, such as generative adversarial networks (GANs) and variational autoencoders (VAEs). Shen et al. (Shen et al., 2019b) extended their original method to make the reconstructions look more natural by conditioning the reconstructed images to be in the subspace of the images generated by a GAN. A similar GAN-prior was used by Yves et al. in (St-Yves & Naselaris, 2018), where the authors also introduced unsupervised training on real-world images. Fang et al. (Fang et al., 2020) leverage the hierarchical structure of the information processing in the visual cortex to propose two decoders, which extract information from the low and high visual cortex areas, respectively. The output of those decoders is used as a conditioning variable in a GAN-based architecture. Shen et al. (Shen et al., 2019a) trained a GAN using a modified loss function that includes an image-space and perceptual loss in addition to the standard adversarial loss. Güçlütürk et al. (Güçlütürk et al., 2017) propose a method to reconstruct perceived faces using a cascade of a linear transformation combined with maximum a posteriori estimation and non-linear transformation combined with adversarial training. A line of work by Seeliger et al. (Seeliger et al., 2018), Mozafari et al. (Mozafari et al., 2020) and Qiao et al. (Qiao et al., 2020) assumes that there exists a linear relationship between the brain activity and the GAN latent space. These methods use the GAN as a real-world image prior to ensure that the reconstructed image has some "naturalness" properties. The work by VanRullen et al. (VanRullen & Reddy, 2019) and Ren et al. (Ren et al., 2021) utilize VAE-GANs (Larsen et al., 2015), a hybrid model in which the VAE decoder and GAN generator are combined. The GAN latent space is used to produce hyperrealistic reconstructions from fMRI activations. The work by Lin et al. (Lin et al., 2022) leverages multi-modality and encodes the fMRI signals into a visual-language latent space and a contrastive loss function to incorporate low-level visual features to the schematic pipeline.

Contributions. In this paper, we propose a novel architecture for the problem of decoding visual imagery from fMRI

recordings. Motivated by the fact that the visual pathway in the human brain processes stimuli in a hierarchical manner, we postulate that such a hierarchy can be captured by the latent space of a deep generative model. More specifically, we use Hierarchical Variational Autoencoders (HVAE) (Vahdat & Kautz, 2020) to learn meaningful representations of stimuli images and we train a deep neural network to learn mappings from the voxel space to the HVAE latent spaces. Voxels originating from the early stages of the visual pathway (V1, V2, V3) are mapped to the earlier layers of latent variables, whereas the higher visual cortex areas (LOC, PPA, FFA) are mapped to the later stages of the latent hierarchy. Our architecture replicates the natural hierarchy of visual information processing in the latent space of a variational model. Our experimental analysis suggests that hierarchical latent models provide better priors for decoding fMRI signals and, to the best of our knowledge, this is the first approach that uses HVAEs in the context of neural decoding.

2. Visual Information Processing

In this section, we give a brief overview of the visual information processing in the human brain and describe the two streams hypothesis, which we use in our experimental architecture. Visual information received from the retina of the eye is interpreted and processed in the visual cortex. The visual cortex is located in the posterior part of the brain, at the occipital lobe, and it is divided into five distinct areas (V1 to V5) depending on the function and structure of the area. Visual stimuli received from the retina travel to the lateral geniculate nucleus (LGN), located near the thalamus. LGN is a multi-layered structure that receives input directly from both retinas and sends axons to the primary visual cortex (V1). V1 is the first and main area of the visual cortex where visual information is received, segmented, and integrated into other regions of the visual cortex. Based on the *two streams hypothesis* (Goodale & Milner, 1992), following V1, visual stimuli can take the *dorsal pathway* or *ventral pathway*. The dorsal pathway consists of the secondary visual cortex (V2), the third visual cortex (V3), and the fifth visual cortex (V5). The dorsal stream, informally known as the "where" stream, is responsible for visually-guided behaviors and localizing objects in space. The ventral stream, also known as the "what" stream, consists of V2 and fourth visual cortex (V4) areas and is responsible for processing information for visual recognition and perception. Visual processing occurs hierarchically at three distinct levels (Groen et al., 2017). The *low-level* includes the retina, lateral geniculate nuclei (LGN), and the primary visual cortex (V1). Low-level processing is the initial step when interpreting an image and it is the place where orientation, edges, and lines are perceived. Sequentially, the *mid-level* processing consists of the secondary (V2), third (V3) and fourth (V4) which extract shapes, ob-

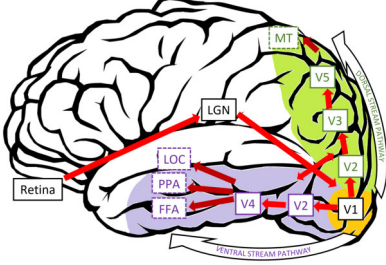


Figure 1. Two stream hypothesis of visual information processing in the human brain.

jects and colors. Finally, the *high-level* processing consists of category-selective areas such as the fusiform face area (FFA), lateral occipital (LOC), parahippocampal area (PPA) and medial temporal area (MT/V5). These areas show selective response to faces, objects/animals, places and motion, respectively. Despite the evident hierarchical structure of visual information processing, most current methods for neural decoding fail to fully exploit that fact. Current methods take into account the hierarchy of visual information processing either by mapping the fMRI voxel to hierarchical CNN-extracted image features via regression models (Shen et al., 2019b; Wen et al., 2018) or by training an end-to-end DNN model on a feature loss function (St-Yves & Naselaris, 2018; Shen et al., 2019a). The major issue with such approaches is that the hierarchy is taken into account in the *feature space* of a CNN model, which is, in general, complex, high-dimensional space. In this work, we propose to take into account the aforementioned hierarchy in the *latent space* of a deep model. Latent spaces are known to produce compact, low-dimensional embeddings of the data and have recently shown impressive performance on image reconstruction and generation tasks (Vahdat & Kautz, 2020). Additionally, the early work of Güçlü et al. in (Güçlü & van Gerven, 2015) and (Güçlü & van Gerven, 2017) reveals a connection between the cortical hierarchy and the hierarchical structure of convolutional neural networks. Given these facts, we postulate that a hierarchical latent space provides better priors for decoding fMRI signals. The intuition is that each brain area, being "responsible" for a certain set of features, better be mapped on a compact, low-dimensional representation of those features. For example, given that V1 is broadly responsible for encoding low-level features (e.g., edges, orientations), it is sensible to map the fMRI voxels from the V1 region onto a representation of the underlying images features; and this mapping is much easier to be learned on the latent space, rather than the feature space.

3. Method

Leveraging the aforementioned intuition, we introduce a neural decoding method that mimics the hierarchical vi-

sual information processing in the *latent space*. Our architecture has two main components: a *Hierarchical Variational Autoencoder (HVAE)* and a *Neural Decoder*. The HVAE is used for learning compact, hierarchical latent representations of real-world images and is trained using self-supervision. The Neural Decoder is used for mapping the brain signals to the HVAE hierarchical latent space and is trained via supervision on {fMRI, Image} pairs. In this section, we describe each of the components in more detail. Our architecture is visualized in Fig 2 for the special case of 2 latent hierarchical layers.

3.1. Hierarchical Variational Autoencoders

To capture the inherent hierarchical structure of visual information processing, we propose to model images via a family of probabilistic models known as Hierarchical Variational Autoencoders (HVAEs). HVAEs extend the basic Variational Autoencoder (VAEs) by introducing a hierarchy of latent variables. Formally, let $\mathbf{x}$ be an image and $\mathbf{z} = \{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_L\}$ be a set of L latent variables. The *generative distribution* or *decoder* is defined as $p_\theta(\mathbf{x}|\mathbf{z}) = p_\theta(\mathbf{x}|\mathbf{z}_1) \prod_{i=1}^{L-1} p_\theta(\mathbf{z}_{i+1}|\mathbf{z}_i)$ and is parametrized by θ . The prior distribution is defined as $p(\mathbf{z}) = p(\mathbf{z}_1) \prod_{i=1}^{L-1} p(\mathbf{z}_{i+1}|\mathbf{z}_i)$. The posterior $p(\mathbf{z}|\mathbf{x})$ is approximated by the *variational distribution* or *encoder* $q_\phi(\mathbf{z}|\mathbf{x}) = q_\phi(\mathbf{z}_1|\mathbf{x}) \prod_{i=1}^L q_\phi(\mathbf{z}_{i+1}|\mathbf{z}_i)$, which is parametrized by ϕ . Both the prior and the approximate posterior are represented by factorial Normal distributions. The variational principle provides a tractable lower bound, known as *Evidence Lower Bound (ELBO)*, on the log-likelihood, as follows

$$\begin{aligned} \log p_\theta(\mathbf{x}) &\geq \mathbb{E}_{q_\phi(\mathbf{z}|\mathbf{x})} \left[\log \frac{p_\theta(\mathbf{x}, \mathbf{z})}{q_\phi(\mathbf{z}|\mathbf{x})} \right] = \mathcal{L}(\theta, \phi; \mathbf{x}) \\ &= -KL(q_\phi(\mathbf{z})||p_\theta(\mathbf{z})) + \mathbb{E}_{q_\phi(\mathbf{z}|\mathbf{x})} [\log p_\theta(\mathbf{x}|\mathbf{z})], \quad (1) \end{aligned}$$

where KL is the Kullback-Leibler divergence. The encoder and decoder are implemented by deep neural networks and their parameters are jointly optimized using gradient descent on the ELBO criterion. Similarly to standard VAEs, the reparametrization trick (Kingma & Welling, 2014; Rezende et al., 2014) is used to allow us to back-propagate the gradient through the stochastic sampling involved in the computation of Eq. 1.

3.2. Neural Decoder

We now leverage the latent space of the HVAE to learn a set of maps from the fMRI voxel space to the hierarchical latent variables. In more detail, each *region of interest (ROI)* is mapped via a dense neural network to a specific subset of the latent space. Brain regions in the earlier states of the visual pathway are mapped to the earlier layers of the latent hierarchy, whereas voxels from the higher visual cortex

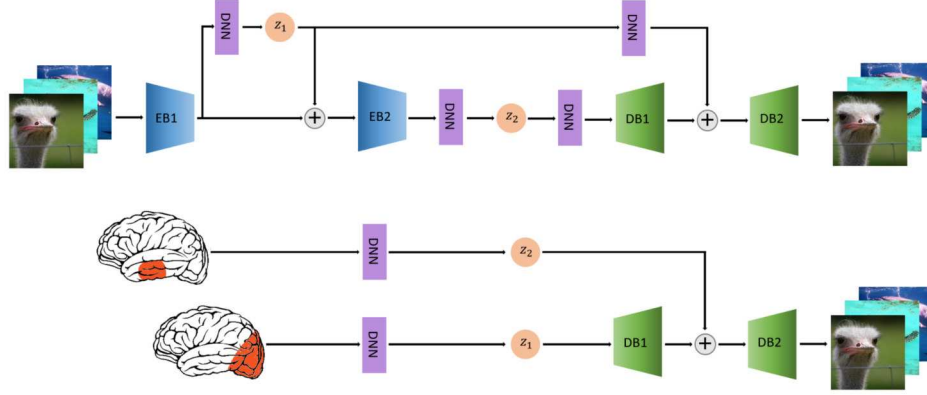


Figure 2. **Outline of our method:** a) We pretrain a Hierarchical Variational Autoencoder on a large set of images. Two layers of latent variables z_1, z_2 are inserted after each encoder (EB) and decoder (DB) block. b) We train the Neural Decoder by discarding the encoder from the previous step and learning a map from the fMRI voxels to the hierarchical latent space. The lower visual cortex (V1, V2, V3) is mapped to z_1 and the higher visual cortex (FFA, PPA, LOC) to z_2 .

areas are mapped to the deep layers in the latent hierarchy. We assume that the HVAE has L groups of latent variables $z_1, z_2, \dots, z_L$ and that the fMRI voxels are partitioned into L non-overlapping brain regions of choice, i.e., $y_1, y_2, \dots, y_L$. Formally, the *Neural Decoder* is a set of maps from the i -th brain region to the i -th group of latent variables. Each of these maps is represented by a deep neural network with parameters w_i , i.e., $z_i = \psi_{w_i}(y_i)$, $i = 1, 2, \dots, L$. The reconstruction $\hat{x}$ is then obtained by passing the latent variables $z = \{\psi_{w_1}(y_1), \psi_{w_2}(y_2), \dots, \psi_{w_L}(y_L)\}$ through the decoder model $p(x|z)$ defined in Sec. 3.1.

The loss function used for training the Neural Decoder is an important design choice. Classic per-pixel measures, such as Euclidean distance, commonly used for regression problems, or the related Peak Signal-to-Noise Ratio (PSNR), are insufficient for images, as they assume pixel-wise independence. Therefore, to encourage the Neural Decoder to learn reconstructions guided by human visual perception, we use a *perceptual loss*. Perceptual loss is a class of loss functions that relies on the fact that CNNs extract hierarchical features. More specifically, deep features trained on supervised, self-supervised and unsupervised objectives are an effective model of human visual perceptual similarity (Zhang et al., 2018). For a given image x and its reconstruction $\hat{x}$, their perceptual loss is:

$$l(x, \hat{x}) = \sum_m \frac{1}{H_m W_m} \sum_{h,w} \|b_m \odot (f_{x,hw}^l - f_{\hat{x},hw}^l)\|_2^2, \quad (2)$$

where $f_{x,hw}^l, f_{\hat{x},hw}^l$ are the layer-wise activations of a given, pretrained CNN model, $b_l \in \mathbb{R}^{C_l}$ is a channel-wise scaling vector. Intuitively, the perceptual loss in Eq. 2 extracts features for both the target and reconstructed image and then compares the features layer-wise using the Euclidean norm.

To ensure that no bias is introduced during learning, it is important that the CNN used for evaluating Eq. 2 is different than the one used for the encoder. In our implementation we use a pretrained AlexNet as well as the code by Zhang et al. (Zhang et al., 2018) to compute the perceptual loss.

3.3. Model Training

For the **encoder** part of our HVAE, we use a pretrained VGG-19 model (Simonyan & Zisserman, 2015). This is a deep convolutional neural network of 19 convolutional layers and 3 fully connected layers. We use the weights from the model pretrained on ImageNet and discard the fully connected layers. We introduce latent variables by taking the output of a given convolutional layer, flattening it, passing it through a fully connected layer and, finally, through a variational layer which outputs the latent variable. This latent variable is re-sampled to avoid any dimension mismatch, and rerouted back to the main block, where it is aggregated with the output of the convolutional layer. Depending on how many latent layers we would like to insert, their exact position may vary. As an empirical design choice we choose to insert the latent layers equally spaced and after a convolutional block. A latent layer is always inserted at the output of the penultimate convolutional block.

The **decoder** part of our HVAE transforms the hierarchical latent variables to output images and consists of 4 transposed convolutional layers. The number of decoder filters are $[128, 64, 32, 16, 3]$ and all kernel sizes are set to 5. Each transposed convolutional layer is followed by a 2d batch normalization and a ReLU non-linearity. The output of each transposed convolutional layer is interleaved with the latent variables. More specifically, each latent variable is initially passed through a fully connected layer, re-sampled to avoid

dimension mismatch and then aggregated with the output of the corresponding transposed convolution. Similarly to the encoder, we insert the latent variable such that we ensure symmetry and we always insert the penultimate latent variable before the first transposed convolution.

We start the training process by first deciding the number and position of the latent layers. The choice is guided by the type of fMRI data that we have as well as the level of latent space coarse-graining that we can achieve. For instance, if our fMRI data contains only the primary (V1) and the secondary (V2) visual cortex then we have two choices: a) we can either consolidate all voxels into a single vector and have a single latent layer in our HVAE or b) we can have two vectors containing the voxels from each brain area and train the HVAE such that it has two latent layers z_1, z_2 (example shown in Fig. 2). Naturally, if our fMRI data are more fine grain, we can add additional latent layers.

Following this design choice, the training proceeds in two phases: In the first phase, we pretrain the HVAE via self-supervision using the ELBO loss function Eq. 1 on a large ensemble of 50,000 real-world images from the ImageNet database. These images come from the same categories as the images shown to the subjects but no test images are included. This phase gives us meaningful latent representations and allows the HVAE decoder to adapt to the statistics of a large set of real-world images. In the second phase, the HVAE encoder is discarded, the HVAE decoder is kept fixed and the Neural Decoder is trained on supervised {fMRI, Image} pairs using the perceptual loss function Eq. 2. In this phase, we essentially learn a map from the voxels of each brain area to the corresponding latent layer and then use that latent vector to reconstruct the image.

4. Experimental Results

To evaluate the utility of our method in practice, we carry out a series of experimental simulations. To measure the performance of our method, we use both qualitative comparisons of the reconstructions as well as quantitative metrics. In what follows, we give the details of the dataset used, the metrics implemented and baseline comparisons.

Dataset: We applied our pipeline on a commonly used, publicly available dataset known as **Generic Object Decoding (GOD)**. The dataset consists of high-resolution (500×500) stimuli images and their corresponding fMRI recordings. There exist 1250 (1200 train, 50 test) stimuli images selected from 200 object categories from the ImageNet database and the fMRI recording were obtained while 5 healthy subjects were viewing the stimuli (presentation experiment). The train- and test-fMRI data consist of 1 and 35 (repeated recordings) per presented stimulus image, respectively. We use the post-processed

fMRI data provided by Horikawa et al. (Horikawa & Kamitani, 2017), which contain voxels from 7 brain areas (V1, V2, V3, V4, FFA, PPA, LOC). The temporal component of the fMRI signal is averaged-out and the input to the model is a high-dimensional voxel vector. Even though there may be more comprehensive datasets, such as the BOLD 5000 (Chang et al., 2019) and the NSD (Allen et al., 2022) datasets (which in fact contain a higher number of more diverse images), we choose to focus on GOD for two primary reasons: 1) the dataset provides post-processed fMRI data, and 2) it has been used in numerous past studies (Beliy et al., 2019; Shen et al., 2019b; Fang et al., 2020). Both of these facts facilitate the easy and fair comparison between different methods.

Ablation Study: We perform an ablation study, with the number of hierarchical layers and, consecutively, the number of brain regions, being the ablated parameter. Motivated by the two stream hypothesis (Sec. 2), we consider the following variants:

1. **Naive Baseline (NB):** We consider only one latent layer z_{NB} and all fMRI voxels are mapped to z_{NB} . There are approximately 5000 voxels in this variant.
2. **Primary-Secondary (PS):** We consider 2 latent layers z_{V1}, z_{V2} and the voxels from V1, V2 are mapped to the corresponding latent layer. There are approximately 1500 voxels.
3. **Dorsal Pathway (DP):** We consider the 3 latent layers z_{V1}, z_{V2}, z_{V3} and voxels from V1, V2, V3 are mapped to the corresponding latent. There are approximately 2500 voxels.
4. **Ventral Pathway (VP):** We consider 4 latent layers $z_{V1}, z_{V2}, z_{V4}, z_{PF}$ and the voxels from V1, V2, V4, {FFA, PPA} are mapped to the corresponding latent layer. The voxels from FFA and PPA are merged to a single area. There are approximately 3300 voxels.

We note that by using different ROIs and/or by combining them to form different latent architectures, it is possible to obtain different ablated variants. We empirically noticed that by including the LOC, either concatenated as part of the latest latent layer of the VP or by creating a new LOC-only latent layer, there was no further performance improvements, only losses in terms of computational cost. Therefore, we restrict our exposition to the aforementioned 4 variants.

Metrics: The reconstruction quality is assessed both subjectively, i.e., by visual inspection of the output test images and comparison with the ground truth, as well as objectively. Our quantitative evaluation relies on metrics that encode the

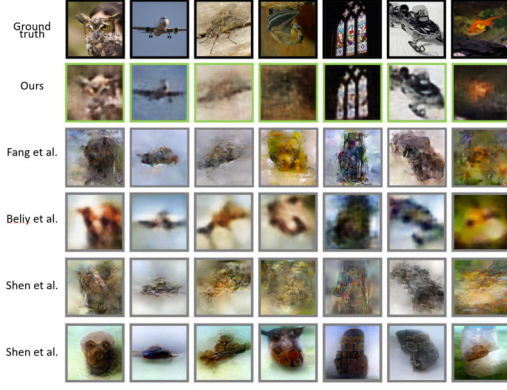


Figure 3. Qualitative comparison of reconstruction quality.

spatial dependence such as the Pearson Correlation Coefficient (PCC) and the Structural Similarity Index Measure (SSIM).

Pearson Correlation Coefficient (PCC): This metric is extensively used in statistics to measure the linear dependence between variables. In the context of image similarity, PCC is computed on the flattened representations of the two images. The limitation of PCC is its sensitivity to edge intensity or misalignment, which makes the metric assign larger value to blurry images (Beliy et al., 2019).

Structural Similarity Index Measure (SSIM): Wang et al. proposed SSIM in (Wang et al., 2004) as a metric that quantifies the characteristics of human vision. Given a pair of images p, q , SSIM is computed as a weighted combination of *luminance*, *contrast* and *structure*. Assuming equal contribution of each measure, SSIM is first computed locally in a common window of size $N \times N$, and then the global SSIM is computed by averaging the SSIM over all non-overlapping windows.

These image similarity metrics defined are used for computing the *correct identification rate* in an **n-way classification** task. Let $M \in \{PCC, SSIM\}$ be a metric of choice, $\hat{p}_i$ be a reconstructed image and P_i be a set containing the ground truth p_i and a set of $n - 1$ randomly selected target images. The *Correct Identification Rate (CIR)* is defined as follows:

$$CIR_M^n = \frac{1}{N} \sum_{i=1}^N \mathbf{1}(i = \arg \max_{p_j \in P_i} M(\hat{p}_i, p_j)), \quad (3)$$

where N is the total number of images and the indicator function $\mathbf{1}(\cdot)$ has the value of 1 if the argument is true and 0 otherwise. The CIR_M^n metric is essentially the frequency at which a reconstructed image can correctly identify the ground truth among $n - 1$ randomly selected additional images. The chance level is $1/n$.

Main Results: We compare the performance of our method

against several state of the art methods (SOTA) for the problem of neural decoding. The competitor methods are: the encoder-decoder based self-supervised method by Beliy et al. (Beliy et al., 2019), the end-to-end, GAN-based pipeline by Shen et al. (Shen et al., 2019a), the GAN-conditioned method by Shen et al. (Shen et al., 2019b) and the shape-schematic GAN by Fang et al. Figure 3 shows qualitative results and compares our method against the aforementioned competitors. All displayed images were reconstructed from the test-fMRI dataset. To improve the signal-to-noise ratio, the test fMRI test samples are averaged across trials. The results shown were obtained using the Ventral Pathway variant, which gave the best performance. We directly use the reconstructions reported in the respective papers by the authors. Our method tends to consistently produce more faithful reconstructions. Note that, even though the GAN-based decoders tend to produce more natural images, the reconstructions may deviate significantly from the stimulus image. This is because the GAN is introduced as an imaged prior, as noted by Beliy et al. (Beliy et al., 2019). On the contrary, our method reconstructs the stimuli more faithfully, albeit the reconstructions appearing as a noisier version of the ground truth. (Fang et al., 2020).

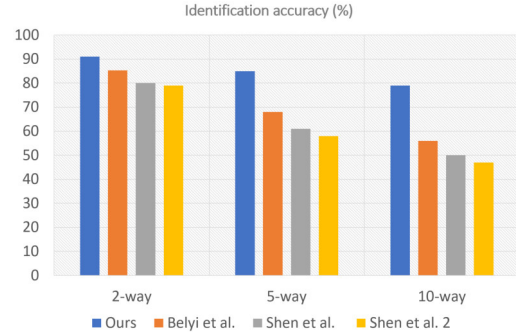


Figure 4. Correct identification ratio.

The qualitative comparison highlights a trade-off between the naturalness of the reconstructed stimuli and the pixel-wise noise introduced in the reconstructions. To resolve the ambiguity, we perform an additional quantitative comparison using the CIR^n metric. For this part we compare against the method by Beliy et al. (Beliy et al., 2019) as well as the two variants of the method by Shen et al. (Shen et al., 2019b). We directly compare against the results as reported by the authors of (Beliy et al., 2019). The results are shown in Fig. 4. For our method, we report the correct identification rate obtained using the Ventral Pathway variant and we average across the metrics (CIR_{PCC}^n and CIR_{SSIM}^n). We observe that our method consistently outperforms the competitors and, particularly in the 5-way and 10-way case, by a substantial margin. Additionally, we observe our method shows a small performance drop as we

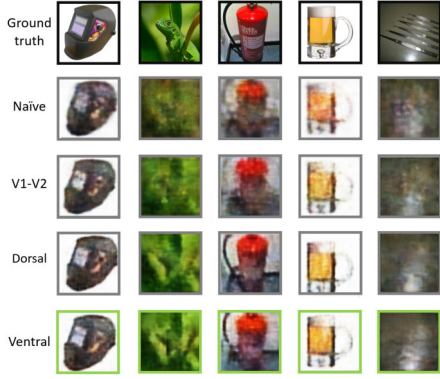


Figure 5. Qualitative comparison for different pathways.

increase n , i.e., from 90% in the 2-way case to 79% in the 10-way case, whereas the performance loss for the competitor method is substantially higher. This performance is due to the following fact: Even though our method gives noisier reconstructions than the competitors, the high-level features such as color, texture and shapes are retained and, therefore, the task of identifying the correct ground truth from the reconstruction is easier. In contrast, please observe in Fig. 3 that the competitor methods may substantially alter the color or texture of the image, therefore leading to more frequent ground truth misidentification.

In the next experiment, we evaluate the decoding performance of different visual pathways. The results are shown on Fig. 5. Qualitatively, the ventral stream seems to be producing the best reconstructions, which is expected from a neuroscience perspective, given that this pathway’s purpose is for visual perception and contains high level areas (FFA-PPA) for object recognition. Interestingly enough, even though the Naive Baseline contains all the available brain areas, the reconstruction quality is inferior, especially in the 2nd and 3rd images, which are far more complex.

The V1-V2, Dorsal and Ventral variants essentially partition the brain areas into (progressively finer) segments and map the voxels from each area onto the hierarchical latent space of the HVAE decoder. Even though the increased performance among these variants may be partially explained by the fact that the number of voxels increases, the main point of comparison should be against the Naive Baseline. The three models, PS, DP, and VP, are hierarchical, whereas the naive baseline includes all data but has no hierarchy. Simply the fMRI responses from two regions, V1 and V2 and discarding all other voxels we are able to achieve better performance than simply mapping all voxels in a big latent vector. This suggests that the hierarchy is far more important than the amount of data that we fed to the model. This is in line with previous studies which concluded that models trained on the whole visual cortex perform slightly worse than those trained on separate areas (Fang et al., 2020).

Additionally, since the Naive Baseline essentially learns a map from *all voxels* to a *single* latent layer, it is natural to assume that it fails due to massively overfitting. However, if overfitting is indeed the only reason for that failure, we would expect the reconstruction performance to decrease as we add more voxels to the model input. However, the figure shows the exact opposite: the performance increases as we add more voxels. This suggests that overfitting is not the only reason for the Naive model’s failure and that the hierarchical structure of the visual information processing needs to be explicitly taken into account. However, one may hypothesize that the performance increase in the Ventral Pathway model may come from the partitioning of the ROIs and that the hierarchical structure has little impact. To test this, it is prudent to include a variant in which the VP ROIs are randomly shuffled to assess whether the hierarchical structure or the partitioning of the voxels drives the performance. We call this variant **VP Permutations** and it supplements the previous 4 variants.

Following that, we present quantitative results on Table 1. On this table we give the n -way correct identification rate CIR^n for $n = 2, 5, 10$, for all ablated variants and the VP Permutations for both metrics (PCC and SSIM). The results on this table validate the aforementioned qualitative observations. The identification accuracy is progressively increasing as we partition the brain into more fine areas and as we add hierarchical layers in the HVAE onto which the brain areas are mapped. Additionally, we observe that the newly introduced VP Permutations variant leads to performance degradation, which suggests that the hierarchy and not the partitioning drives the performance result. However, we do note that we have a slight performance increase compared to the Naive Baseline, which indicates that merely partitioning the brain regions is beneficial, albeit not as beneficial as accounting for the hierarchy.

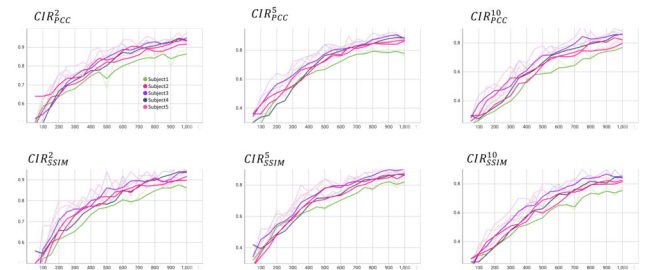


Figure 6. Learning curves for CIR_M^n , $n = 2, 5, 10$ and $M \in \{PCC, SSIM\}$ across all subjects. The horizontal axis is the number of epochs. Subject 3 is marginally outperforming the other subjects and Subject 1 gives the worst performance (figure best viewed in color).

To verify that our method can be successfully applied to all subjects and study potential inter-subject variation of

	CIR_{PCC}^2	CIR_{SSIM}^2	CIR_{PCC}^5	CIR_{SSIM}^5	CIR_{PCC}^{10}	CIR_{SSIM}^{10}
Naive Baseline	0.77	0.78	0.64	0.66	0.57	0.58
Primary-Secondary	0.80	0.82	0.72	0.73	0.65	0.67
Dorsal Pathway	0.88	0.90	0.81	0.80	0.75	0.75
Ventral Pathway	0.91	0.92	0.84	0.85	0.79	0.79
VP Permutations	0.79	0.80	0.65	0.66	0.60	0.58

Table 1. The n -way correct identification rate ($n = 2, 5, 10$) for all ablated variants using the Pearson Correlation Coefficient (PCC) and the Structural Similarity Index Measure (SSIM) as a selection criterion. We report the mean across subjects. The results for VP Permutations are averaged across 4 permutations. The inter-subject deviation was in the range of 0.02 – 0.05. The chance levels are 0.5, 0.25, 0.10, respectively.

the results, we show in Fig. 6 the learning curves for all 5 subjects and for all metrics CIR_M^n , $n = 2, 5, 10$ and $M \in \{PCC, SSIM\}$. The metrics were calculated using the test samples and the ventral pathway variant. Even though the metrics appear similar across subject, after careful examination of the curves some subtle discrepancies and trends can be observed. Subject 1 is consistently performing approximately 5% worse across all metrics whereas Subject 3 is marginally outperforming the other subjects by 2%. The fact that the Subject 3 gives the best reconstructions has been verified in previous studies (Gaziv et al., 2020) and is attributed to differences in the signal-to-noise ratio across subjects. Finally, Fig. 6 allows us to study how training progresses and validate that no overfitting occurs. We observe that, in all cases, the metrics saturate at about 800 epochs, which gives us an empirical estimate of how many iterations our model needs to achieve good performance.

5. Conclusion

We addressed the problem of neural decoding from fMRI recordings and proposed a novel architecture inspired by neuroscience. More specifically, motivated by the fact that the human brain processes visual stimuli in a hierarchical fashion, we postulated that this structure can be captured by latent space of a hierarchical variational autoencoder (HVAE). Our HVAE serves as a proxy to learning meaningful latent representations of stimuli images and can be pretrained on a large dataset of high-resolution images. Following that, we train our Neural Decoder to learn a map from the fMRI voxel space to the HVAE latent space. Our architecture replicates the visual information processing in the human brain in the sense that earlier visual cortex areas (e.g., primary-secondary visual cortex) are mapped to the earlier latent layers, whereas voxels from the higher visual cortex (e.g., PPA, FFA areas) are mapped to the later latent layers. We validated our approach using fMRI recordings from a visual presentation experiment involving 5 subjects and compared against other methods. Our work paves the way to constructing better models to replicate human perception and understanding the nuances of human visual

reconstruction, both of which could be utilized to better understand the brain, assist people with visual disabilities and perhaps in decoding imagery during sleep.

Acknowledgements

P.K., A.I., and P.B. gratefully acknowledge the support by the National Science Foundation under the Career Award CPS/CNS-1453860, the NSF award under Grant Numbers CCF-1837131, MCB-1936775, CMMI-1936624, and CNS-1932620, the U.S. Army Research Office (ARO) under Grant No. W911NF-17-1-0076 and W911NF2310111, and the DARPA Young Faculty Award and DARPA Director Award, under Grant Number N66001-17-1-4044, a National Institutes of Health (NIH) grant 1R01 AG079957, a 2021 USC Stevens Center Technology Advancement Grant (TAG) award, an Intel faculty award and a Northrop Grumman grant. The funder had no role in study design, data collection and analysis, decision to publish, or preparation of the manuscript. The views, opinions, and/or findings contained in this article are those of the authors and should not be interpreted as representing official views or policies, either expressed or implied by the Defense Advanced Research Projects Agency, the Department of Defense, the National Institutes of Health or the National Science Foundation.

References

- Allen, E. J., St-Yves, G., Wu, Y., Breedlove, J. L., Prince, J. S., Dowdle, L. T., Nau, M., Caron, B., Pestilli, F., Charest, I., Hutchinson, J. B., Naselaris, T., and Kay, K. A massive 7t fmri dataset to bridge cognitive neuroscience and artificial intelligence. *Nature Neuroscience*, 25(1): 116–126, Jan 2022. ISSN 1546-1726. doi: 10.1038/s41593-021-00962-x. URL <https://doi.org/10.1038/s41593-021-00962-x>.
- Beliy, R., Gaziv, G., Hoogi, A., Strappini, F., Golan, T., and Irani, M. From voxels to pixels and back: Self-supervision in natural-image reconstruction from fmri. 2019. doi: 10.48550/ARXIV.1907.02431. URL <https://arxiv.org/abs/1907.02431>.
- Brandman, D. M., Cash, S. S., and Hochberg, L. R. Review: Human intracortical recording and neural decoding for brain-computer interfaces. *IEEE Trans. Neural Syst. Rehabil. Eng.*, 25(10):1687–1696, October 2017.
- Chang, N., Pyles, J. A., Marcus, A., Gupta, A., Tarr, M. J., and Aminoff, E. M. Bold5000, a public fmri dataset while viewing 5000 visual images. *Scientific Data*, 6(1):49, May 2019. ISSN 2052-4463. doi: 10.1038/s41597-019-0052-3. URL <https://doi.org/10.1038/s41597-019-0052-3>.
- Fang, T., Qi, Y., and Pan, G. Reconstructing perceptive images from brain activity by shape-semantic gan. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 13038–13048. Curran Associates, Inc., 2020. URL <https://proceedings.neurips.cc/paper/2020/file/9813b270ed0288e7c0388f0fd4ec68f5-Paper.pdf>.
- Gaziv, G., Beliy, R., Granot, N., Hoogi, A., Strappini, F., Golan, T., and Irani, M. Self-supervised natural image reconstruction and rich semantic classification from brain activity. *bioRxiv*, 2020. doi: 10.1101/2020.09.06.284794. URL <https://www.biorxiv.org/content/early/2020/09/08/2020.09.06.284794>.
- Glaser, J. I., Benjamin, A. S., Chowdhury, R. H., Perich, M. G., Miller, L. E., and Kording, K. P. Machine learning for neural decoding. *eNeuro*, 7(4):ENEURO.0506–19.2020, August 2020.
- Goodale, M. A. and Milner, A. D. Separate visual pathways for perception and action. *Trends Neurosci.*, 15(1):20–25, January 1992.
- Groen, I. I. A., Silson, E. H., and Baker, C. I. Contributions of low- and high-level properties to neural processing of visual scenes in the human brain. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, 372(1714), February 2017.
- Güçlütürk, Y., Güçlü, U., Seeliger, K., Bosch, S., van Lier, R., and van Gerven, M. A. J. Reconstructing perceived faces from brain activations with deep adversarial neural decoding. In Guyon, I., Luxburg, U. V., Bengio, S., Wallach, H., Fergus, R., Vishwanathan, S., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL https://proceedings.neurips.cc/paper_files/paper/2017/file/efdf562ce2fb0ad460fd8e9d33e57f57-Paper.pdf.
- Güçlü, U. and van Gerven, M. A. J. Deep neural networks reveal a gradient in the complexity of neural representations across the ventral stream. *J. Neurosci.*, 35(27): 10005–10014, July 2015.
- Güçlü, U. and van Gerven, M. A. J. Increasingly complex representations of natural movies across the dorsal stream are shared between subjects. *Neuroimage*, 145(Pt B): 329–336, January 2017.
- Horikawa, T. and Kamitani, Y. Generic decoding of seen and imagined objects using hierarchical visual features. *Nature Communications*, 8(1):15037, May 2017. ISSN 2041-1723. doi: 10.1038/ncomms15037. URL <https://doi.org/10.1038/ncomms15037>.
- Horikawa, T., Tamaki, M., Miyawaki, Y., and Kamitani, Y. Neural decoding of visual imagery during sleep. *Science (New York, N.Y.)*, 340, 04 2013. doi: 10.1126/science.1234330.
- Kay, K. N., Naselaris, T., Prenger, R. J., and Gallant, J. L. Identifying natural images from human brain activity. *Nature*, 452(7185):352–355, Mar 2008. ISSN 1476-4687. doi: 10.1038/nature06713. URL <https://doi.org/10.1038/nature06713>.
- Kingma, D. P. and Welling, M. Auto-Encoding Variational Bayes. In *2nd International Conference on Learning Representations, ICLR 2014, Banff, AB, Canada, April 14-16, 2014, Conference Track Proceedings*, 2014.
- Larsen, A. B. L., Sønderby, S. K., and Winther, O. Autoencoding beyond pixels using a learned similarity metric. *CoRR*, abs/1512.09300, 2015. URL <http://arxiv.org/abs/1512.09300>.
- Le, L., Ambrogioni, L., Seeliger, K., Güçlütürk, Y., van Gerven, M., and Güçlü, U. Brain2pix: Fully convolutional naturalistic video frame reconstruction from brain activity. *Frontiers in Neuroscience*, 16, 2022. ISSN 1662-453X. doi: 10.3389/fnins.2022.

940972. URL <https://www.frontiersin.org/articles/10.3389/fnins.2022.940972>.
- Lin, S., Sprague, T., and Singh, A. K. Mind reader: Reconstructing complex images from brain activities. *arXiv preprint arXiv:2210.01769*, 2022.
- Miyawaki, Y., Uchida, H., Yamashita, O., Sato, M.-A., Morito, Y., Tanabe, H. C., Sadato, N., and Kamitani, Y. Visual image reconstruction from human brain activity using a combination of multiscale local image decoders. *Neuron*, 60(5):915–929, December 2008.
- Mozafari, M., Reddy, L., and VanRullen, R. Reconstructing natural scenes from fmri patterns using bigbigan. *CoRR*, abs/2001.11761, 2020. URL <https://arxiv.org/abs/2001.11761>.
- Qiao, K., Chen, J., Wang, L., Zhang, C., Tong, L., and Yan, B. Biggan-based bayesian reconstruction of natural images from human brain activity. *CoRR*, abs/2003.06105, 2020. URL <https://arxiv.org/abs/2003.06105>.
- Ren, Z., Li, J., Xue, X., Li, X., Yang, F., Jiao, Z., and Gao, X. Reconstructing seen image from brain activity by visually-guided cognitive representation and adversarial learning. *NeuroImage*, 228:117602, 2021. ISSN 1053-8119. doi: <https://doi.org/10.1016/j.neuroimage.2020.117602>. URL <https://www.sciencedirect.com/science/article/pii/S1053811920310879>.
- Rezende, D. J., Mohamed, S., and Wierstra, D. Stochastic backpropagation and approximate inference in deep generative models. In Xing, E. P. and Jebara, T. (eds.), *Proceedings of the 31st International Conference on Machine Learning*, volume 32 of *Proceedings of Machine Learning Research*, pp. 1278–1286, Beijing, China, 22–24 Jun 2014. PMLR. URL <https://proceedings.mlr.press/v32/rezende14.html>.
- Seeliger, K., Güçlü, U., Ambrogioni, L., Güçlütürk, Y., and van Gerven, M. Generative adversarial networks for reconstructing natural images from brain activity. *NeuroImage*, 181:775–785, 2018. ISSN 1053-8119. doi: <https://doi.org/10.1016/j.neuroimage.2018.07.043>. URL <https://www.sciencedirect.com/science/article/pii/S105381191830658X>.
- Shen, G., Dwivedi, K., Majima, K., Horikawa, T., and Kamitani, Y. End-to-end deep image reconstruction from human brain activity. *Frontiers in Computational Neuroscience*, 13, 2019a. ISSN 1662-5188. doi: 10.3389/fncom.2019.00021. URL <https://www.frontiersin.org/article/10.3389/fncom.2019.00021>.
- Shen, G., Horikawa, T., Majima, K., and Kamitani, Y. Deep image reconstruction from human brain activity. *PLOS Computational Biology*, 15(1):1–23, 01 2019b. doi: 10.1371/journal.pcbi.1006633. URL <https://doi.org/10.1371/journal.pcbi.1006633>.
- Simonyan, K. and Zisserman, A. Very deep convolutional networks for large-scale image recognition. In Bengio, Y. and LeCun, Y. (eds.), *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, 2015. URL <http://arxiv.org/abs/1409.1556>.
- St-Yves, G. and Naselaris, T. Generative adversarial networks conditioned on brain activity reconstruct seen images. In *2018 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*, pp. 1054–1061, 2018. doi: 10.1109/SMC.2018.00187.
- Vahdat, A. and Kautz, J. Nvae: A deep hierarchical variational autoencoder, 2020. URL <https://arxiv.org/abs/2007.03898>.
- VanRullen, R. and Reddy, L. Reconstructing faces from fmri patterns using deep generative neural networks. *Communications Biology*, 2(1):193, May 2019. ISSN 2399-3642. doi: 10.1038/s42003-019-0438-y. URL <https://doi.org/10.1038/s42003-019-0438-y>.
- Wang, Z., Bovik, A. C., Sheikh, H. R., and Simoncelli, E. P. Image quality assessment: from error visibility to structural similarity. *IEEE transactions on image processing*, 13(4):600–612, 2004.
- Wen, H., Shi, J., Zhang, Y., Lu, K.-H., Cao, J., and Liu, Z. Neural encoding and decoding with deep learning for dynamic natural vision. *Cereb. Cortex*, 28(12):4136–4160, December 2018.
- Yoshida, T. and Ohki, K. Natural images are reliably represented by sparse and variable populations of neurons in visual cortex. *Nature Communications*, 11(1):872, Feb 2020. ISSN 2041-1723. doi: 10.1038/s41467-020-14645-x. URL <https://doi.org/10.1038/s41467-020-14645-x>.
- Zhang, R., Isola, P., Efros, A. A., Shechtman, E., and Wang, O. The unreasonable effectiveness of deep features as a perceptual metric. *CoRR*, abs/1801.03924, 2018. URL <http://arxiv.org/abs/1801.03924>.